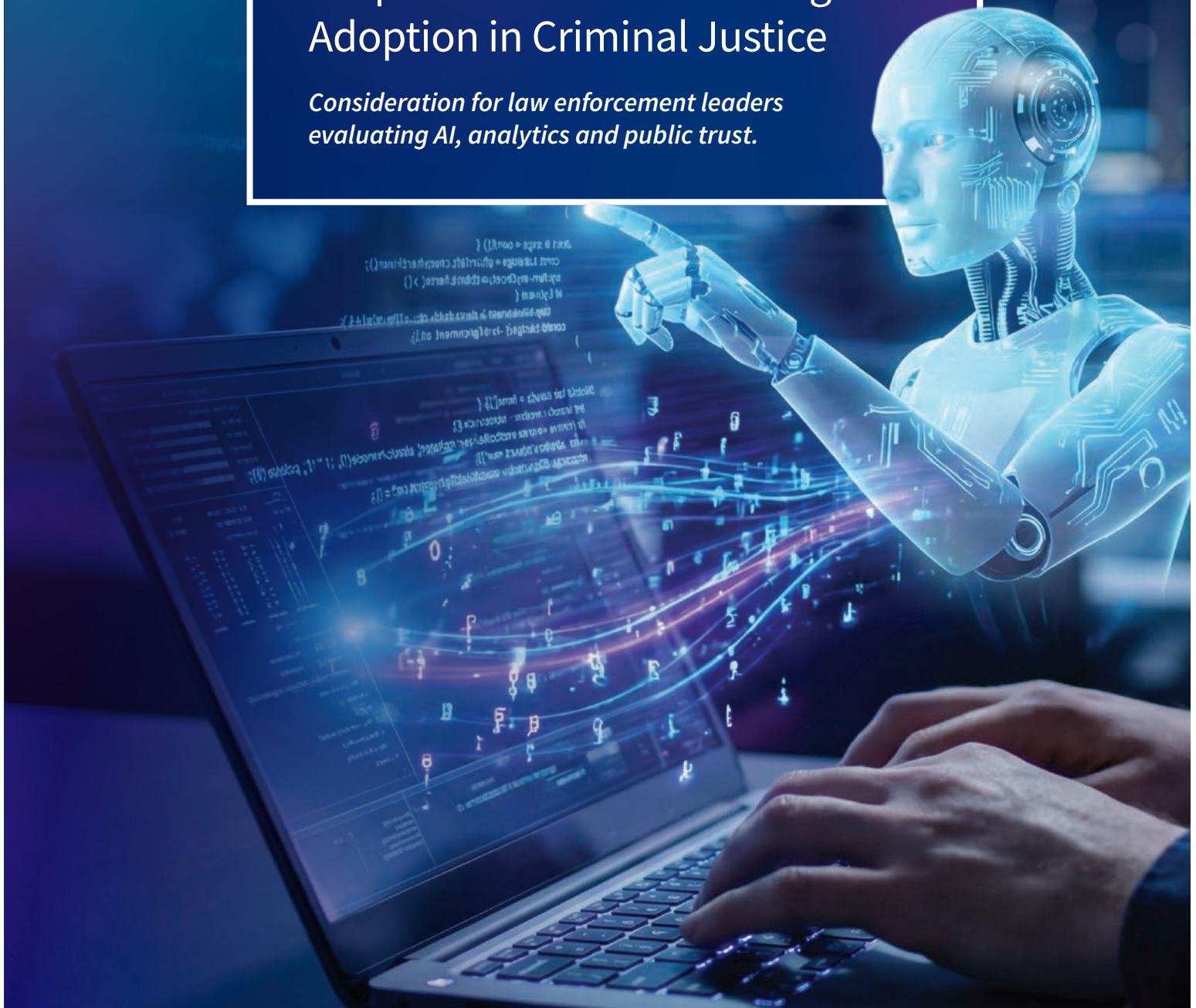


ARTICLE



A Practical Framework for Responsible Artificial Intelligence Adoption in Criminal Justice

Consideration for law enforcement leaders evaluating AI, analytics and public trust.



Criminal justice agencies stand at a critical inflection point. As data volumes grow, investigative timelines compress and operational demands often continue to outpace staffing. Artificial intelligence (AI) is no longer simply an emerging technology. Instead, AI is becoming part of the information infrastructure that will shape how public safety professionals organize facts, identify relationships and make sense of increasingly complex challenges.

For law enforcement leaders, the question is not whether AI will influence criminal justice. AI already is. The more important question is how agencies will adopt AI in a way that strengthens mission outcomes while preserving the trust, accountability and professional judgment that define responsible public safety work.

AI can help agencies work more efficiently, uncover meaningful connections, and make better use of available information. But criminal justice decisions directly affect people, communities, officer safety and public confidence. Any AI-enabled solution used in this environment must therefore be evaluated not only for what it can do, but for how it is developed, governed and applied, including its intended and unintended effects.

Responsible AI adoption requires a leadership framework. Responsible adoption requires clear operational purpose, disciplined governance, transparency, accountability, privacy safeguards and a commitment to keeping consequential decisions in the hands of trained professionals. These guiding principles can help public safety leaders evaluate AI-enabled solutions with integrity, precision and purpose.



EXECUTIVE SUMMARY

Artificial intelligence is becoming part of the information infrastructure that supports modern criminal justice operations. While AI can help agencies organize information, identify meaningful relationships and improve efficiency, AI adoption must be guided by principles that preserve public trust.

Key considerations for law enforcement leaders:



TRANSPARENCY



RESPONSIBLE USE



UNDERSTANDABILITY



SUBSTANTIATION



TRACEABILITY

Together, these principles can help ensure that AI strengthens decision-making while preserving accountability, privacy and professional judgment. These principles also align with emerging guidance from the National Institute of Standards and Technology (NIST), the Council on Criminal Justice, and the U.S. Department of Justice regarding trustworthy AI and effective governance.

BOTTOM LINE: The question is not whether AI can be used, but whether it can be used responsibly, transparently and in a manner worthy of the trust placed in public safety agencies.

Make AI transparent

Criminal justice agencies need technology they can trust, and that trust begins with visibility. Leaders should understand not only what an AI-enabled capability can do, but the purpose it is designed to serve, the information it uses, the limits of its output and the circumstances in which it should not be relied upon. Transparency helps provide agencies the context needed to decide whether a capability belongs in a particular workflow and how it should be governed.



Before applying AI or advanced analytics to a public safety challenge, leaders should examine the problem the technology is intended to help solve. **Who will use it? What data will it rely on? What assumptions may shape the result? How could the output influence an investigation, analysis or operational decision?** These questions help agencies distinguish a mission-aligned capability from technology adopted without a clearly defined role.

Transparency must also reflect the needs of different users. Investigators may need to understand how a tool surfaces connections among people, locations or assets. Analysts may need visibility into how information is searched, organized or prioritized. Agency leaders may require insight into oversight responsibilities, safeguards and appropriate-use boundaries. The goal is not simply to disclose technical detail, but to provide the right information so each user can apply the technology with greater confidence and control.

Put responsible use into practice

Public safety technology carries significant responsibility because its use can affect people, communities, officer safety and public confidence. Responsible AI requires more than strong technical performance. AI requires thoughtful consideration of how a capability may be used, the environment in which it will operate, and the consequences that could follow from its deployment.

For law enforcement leaders, responsible use should be treated as an operating discipline rather than a procurement milestone. Agencies need clear policy, role-based access, training, supervision and review before an AI-enabled capability can become part of routine work. Leaders should also define where technology can provide assistance and where professional judgment, independent verification or additional authorization is required.

AI-enabled solutions should support lawful, ethical and appropriate public safety outcomes. They may organize information, identify patterns or support investigative workflows, but authorized professionals and their agencies should remain responsible for interpreting the information and determining what action, if any, should follow.

The standard should be clear: technology can assist the work, but it cannot assume responsibility for the decision.

Make AI understandable

AI is most useful when professionals can understand the information it provides and apply it appropriately. An output that appears authoritative but cannot be meaningfully interpreted can introduce uncertainty rather than reduce it. Leaders should expect AI-enabled insights to be presented in ways that help users understand what a result means, what it does not mean, and what additional review may be necessary.

Explanation needs vary by role. An investigator may need to understand how a potential relationship was identified. An analyst may want to know how information was organized or prioritized. Command staff may need insight into governance, oversight and risk considerations. Explainability should therefore be proportionate to the capability, the operating context and the responsibilities of the person relying on it.

If a tool highlights a possible lead, users should be able to review the information that contributed to that result, rather than rely solely on the output itself. Making AI understandable keeps trained professionals at the center of decision-making, helps them recognize limitations and supports more confident, defensible use of advanced technology.



Require substantiated outcomes

Confidence in AI depends on confidence in the information and practices behind it. AI-enabled outcomes should be grounded in reliable data, rigorous development methods and appropriate testing. Accuracy, quality and fairness do not occur automatically; they require continual attention throughout design, development, deployment and use.

Unintended bias can enter through data inputs, model development or operational implementation. Agencies should ask how data quality is assessed, how assumptions are documented, what testing is performed and how performance is monitored over time. Agencies should also understand the limitations of source information and the safeguards used to identify issues, before they influence operational outcomes.

Substantiation helps leaders distinguish practical AI from technology driven primarily by hype or speculation. Insights should be supported by evidence that authorized users can review and, when appropriate, independently verify.

Keep AI traceable

In criminal justice, decisions should not depend on conclusions that cannot be examined. Traceability allows users to connect an AI-supported insight to the information, sources and processes that contributed to it. That visibility supports accountability, oversight and confidence in the result.

Traceability also reinforces human responsibility. Investigators, analysts, command staff, records personnel and other authorized users must be able to review supporting information, apply their training, and experience and exercise independent judgment. AI may help organize information, reveal patterns and streamline workflows, but people should remain responsible for deciding what the information means and what action, if any, is appropriate.

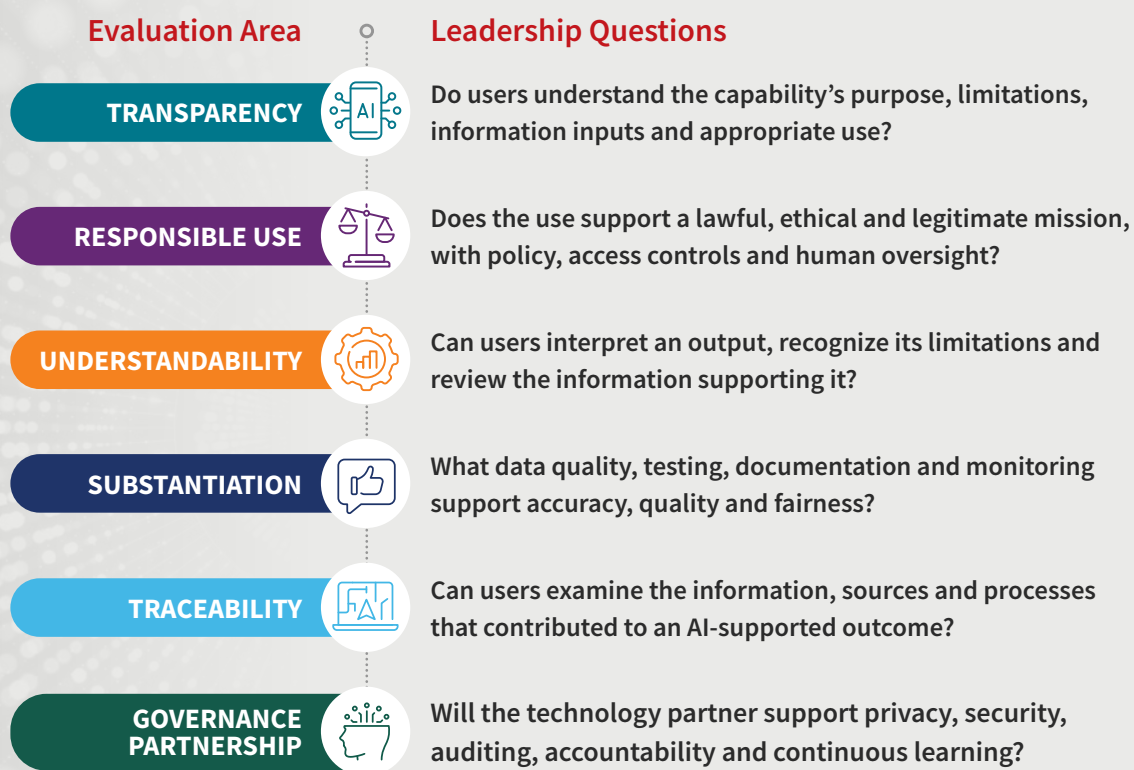
The same principle strengthens data governance. Leaders should understand where information comes from, how it is collected and protected, who can access it, and what auditing and monitoring controls apply. When agencies can trace both the source of an insight and the way it was used, they are better positioned to protect sensitive information, demonstrate responsible stewardship and keep consequential decisions accountable to people, not algorithms.



Leadership Lens: The value of AI in public safety should be measured not by the volume of information it can process, but by the integrity, reliability and public benefit of the insights it supports. Innovation earns legitimacy when stewardship, transparency, accountability, privacy and security advance alongside capability.

A practical evaluation checklist for agency leaders

Before adopting or expanding an AI-enabled capability, leaders can use the following questions to help guide vendor discussions, internal governance reviews and policy alignment.



Building trust through responsible use

Trust is the condition that allows advanced technology to create public value in criminal justice. Agencies can deploy more capable tools, connect more information and accelerate more workflows, but the legitimacy of those capabilities depends on whether leaders can explain how those capabilities are governed, how outputs are verified, and how people remain accountable for the decisions that follow.

For law enforcement leaders, responsible AI adoption should therefore be treated as a leadership discipline, not a procurement milestone. AI adoption requires clear policy, role-based access, training, auditability and a shared understanding of when AI-generated or AI-assisted information may be used. This process also requires a culture that encourages users to question outputs, review supporting information and document the basis for action.

Responsible AI adoption is especially important as public safety information becomes more connected. When analytics can organize complex information, reveal meaningful relationships and provide greater context, the operational benefit can be significant. But greater context must be matched by greater care. Agencies should be able to show that sensitive information is protected, insights are substantiated and outcomes can be traced to the information that supports them.

Responsible use can also strengthen community confidence. Public trust is not built through capability alone. It is built through transparency, responsible stewardship, understandable results, reliable evidence, traceability and consistent adherence to agency policy and applicable law. These practices help demonstrate that AI supports legitimate public safety missions, rather than replacing professional judgment or extending technology beyond its appropriate role.

AI will continue to evolve, and its role in criminal justice and public safety will likely grow with it. Agencies should welcome innovations that can streamline work, improve effectiveness and help professionals respond to increasingly complex challenges. But those innovations must be adopted thoughtfully, with governance equal to the power of the technology itself.

By emphasizing transparency, responsible use, understandability, substantiation and traceability, agencies can benefit from AI while keeping privacy, fairness and human judgment at the center of consequential decisions. The future of responsible AI in criminal justice will be shaped not only by what technology can accomplish, but by the confidence law enforcement leaders can build in how it is developed, governed and used.

The emerging consensus on responsible AI

As artificial intelligence becomes increasingly integrated into public safety operations, a consensus is beginning to emerge around the principles needed to responsibly govern AI use. While organizations may use different terminology, many of the same themes appear repeatedly across government guidance, industry frameworks, and public-sector discussions about trustworthy AI. These common themes reflect a growing recognition that technological advancement alone is not enough.

Trust must be intentionally built into the way AI-enabled capabilities are developed, deployed and governed.

Transparency, human oversight, accountability, explainability, data quality, privacy protection and traceability have become foundational concepts in conversations about responsible AI adoption. The National Institute of Standards and Technology describes trustworthy AI in terms that include validity and reliability, safety, security and resilience, accountability and transparency, explainability and interpretability, privacy enhancement and harmful-bias management.¹ The Council on Criminal Justice Task Force on Artificial Intelligence has also released a framework to help guide the responsible integration of AI across law enforcement, courts and corrections that includes principles of safety, security, effectiveness, fairness, justness and accountability.²

This trend is particularly significant for criminal justice agencies. Unlike many commercial environments, public safety organizations operate in contexts where decisions may directly affect people, communities, officer safety and public confidence. As a result, expectations surrounding AI governance are often higher. The U.S. Department of Justice's Artificial Intelligence and Criminal Justice Final Report addresses the use of AI in criminal justice and is part of the Department's published AI resources.³ For law enforcement leaders, the practical implication is clear: agencies should expect AI-enabled solutions to be understandable, appropriately governed, supported by reliable information and subject to meaningful human oversight. Agencies increasingly seek assurance that solutions can be understood, that outputs can be examined and verified, that sensitive information is protected and that trained professionals remain accountable for consequential decisions.

The principles discussed throughout this paper reflect that broader evolution in thinking. Transparency helps users understand how capabilities function and where they should be applied. Responsible use helps ensure that technology remains aligned with lawful and ethical public safety missions. Understandability enables professionals to evaluate AI-assisted information in context. Substantiation reinforces confidence through reliable data and rigorous development practices. Traceability supports accountability by connecting analytical outputs to the information and processes that informed them. Together, these concepts provide a practical foundation for evaluating whether an AI-enabled solution is worthy of the trust placed in it.



Ultimately, trust is not created through algorithms alone. Instead, trust is earned through stewardship, governance, transparency, accountability and responsible implementation. As public safety agencies continue evaluating emerging technologies, leaders will increasingly focus not only on what AI can do, but on the principles that govern how it is used. That shift from capability alone to capability guided by trust may prove to be one of the most important developments in the future of responsible AI adoption.

Responsible innovation and the Accurint One™ standard

As organizations across government and industry continue to emphasize transparency, accountability, human oversight and responsible governance in AI, Accurint One reflects the belief that public safety analytics must be both powerful and principled. Accurint One represents the practical application of those principles within a public safety environment.

For more than thirty-five years, LexisNexis® Risk Solutions has helped criminal justice professionals transform information into insight. Accurint One represents the next evolution of that mission by combining responsibly sourced data, advanced analytics and artificial intelligence to help authorized professionals better understand the people, places and events they can encounter while preserving human judgment in every consequential decision.

The objective is not simply to adopt more advanced technology. Instead, the objective is to adopt technology worthy of the trust placed in it.

For law enforcement leaders evaluating AI, this distinction matters. The objective is not simply to adopt more advanced technology. Instead, the objective is to adopt technology worthy of the trust placed in it by agencies, officers and the communities they serve. That trust is earned when capabilities are transparent, use is responsible, insights are understandable, outcomes are substantiated and the information behind them can be traced.

The Accurint One Standard reinforces those expectations through lawful data sourcing, rigorous quality standards, disciplined governance, transparent provenance and layered safeguards that include controlled access, auditing, monitoring and accountability. The AI capabilities in Accurint One are designed to support authorized professionals conducting legitimate public safety missions within the legal, ethical, and operational frameworks that govern their agencies.

The role of Accurint One is to provide accurate information, meaningful analytical capabilities and responsible innovation so professionals can clearly and efficiently understand complex situations and make better-informed decisions. The exercise of judgment remains where it belongs: with the trained professionals entrusted to serve their communities.

As AI becomes more capable and information becomes more connected, responsible adoption will require enduring partnership between technology developers and public safety agencies. Innovation, stewardship and professional judgment must advance together. Those principles are the standard law enforcement leaders should demand from any AI-enabled solution and are the standard on which Accurint One is built.



Sources

1. NIST Artificial Intelligence Risk Management Framework 1.0, <https://www.nist.gov/itl/ai-risk-management-framework>
2. Council on Criminal Justice Principles for the Use of AI in Criminal Justice, <https://counciloncj.org/principles-for-the-use-of-ai-in-criminal-justice/>
3. U.S. Department of Justice Artificial Intelligence and Criminal Justice Final Report, <https://www.justice.gov/olp/media/1381796/dl>

**For more information on Accurant One™ and our
Responsible Artificial Intelligence Principles for Criminal Justice Solutions,
please visit risk.lexisnexis.com/AccurantOne.**



About LexisNexis Risk Solutions

LexisNexis® Risk Solutions provides customers with information-based analytics and decision tools that combine public and industry-specific content with advanced technology and algorithms to assist them in evaluating and predicting risk and enhancing operational efficiency. Headquartered in metro Atlanta, Georgia, the company has offices throughout the world, serves customers in more than 190 countries and territories and is part of RELX. For more information, please visit LexisNexis Risk Solutions.

This white paper is provided solely for general informational purposes and presents only summary discussions of the topics discussed. This white paper does not represent legal advice as to any factual situation; nor does it represent an undertaking to keep readers advised of all relevant developments. Readers should consult their attorneys, compliance departments and other professional advisors about any questions they may have as to the subject matter of this white paper.

Accurant One services are not provided by "consumer reporting agencies," as that term is defined in the Fair Credit Reporting Act (15 U.S.C. § 1681, et seq.) ("FCRA") and do not constitute "consumer reports," as that term is defined in the FCRA. Accordingly, Accurant One services may not be used in whole or in part as a factor in determining eligibility for credit, insurance, employment, or another purpose in connection with which a consumer report may be used under the FCRA. Due to the nature of the origin of public record information, the public records and commercially available data sources used in reports may contain errors. Source data is sometimes reported or entered inaccurately, processed poorly or incorrectly, and is generally not free from defect. These products or services aggregate and report data, as provided by the public records and commercially available data sources, and are not the source of the data, nor are they a comprehensive compilation of the data. Before relying on any data, it should be independently verified. LexisNexis and the Knowledge Burst logo are registered trademarks of RELX Inc. Accurant is a registered trademark of LexisNexis Risk Data Management Inc. Accurant One is a registered trademark of LexisNexis Risk Solutions FL Inc. Other products and services may be registered trademarks or trademarks of their respective companies © 2026 LexisNexis Risk Solutions. NXR17147-00-0826-EN-US